



THE UNIVERSITY *of* EDINBURGH
informatics

Applied Machine Learning (AML)

Non-Linear Dimensionality Reduction

Oisín Mac Aodha • Siddharth N.

Outline

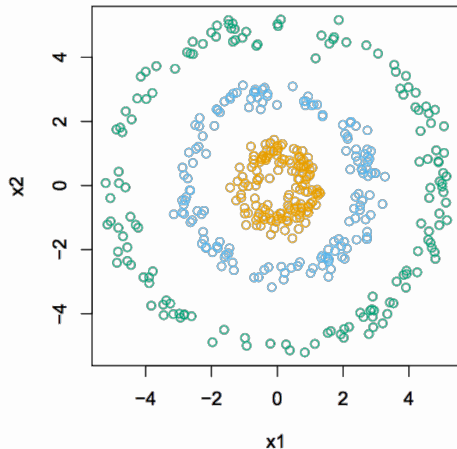
Non-Linearity

- Extensions for Linear Dimensionality Reduction
 - Kernel PCA
- Visualisation
 - Multi-Dimensional Scaling (MDS)
 - Isomap
 - Locally Linear Embeddings (LLE)
 - t -distributed Stochastic Neighbour Embedding (t -SNE)
 - Uniform Manifold Approximation & Projection (UMAP)

Extensions for Linear Dimensionality Reduction

PCA on Non-Linear Data

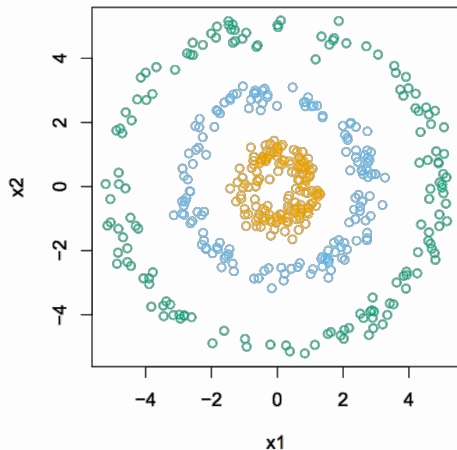
Example: Shells



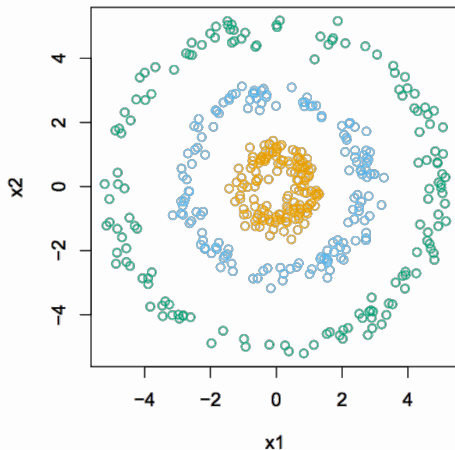
Data

PCA on Non-Linear Data

Example: Shells



Data



PCA

Feature Transformation

Key Idea: Transform inputs x using a feature map $\phi(x)$

Feature Transformation

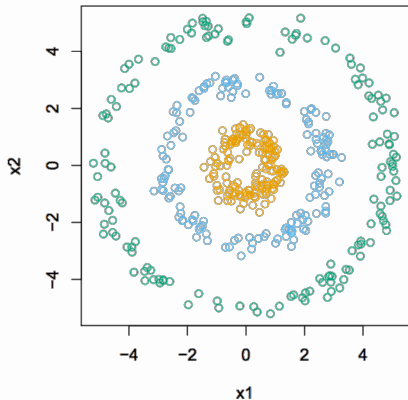
Key Idea: Transform inputs x using a feature map $\phi(x)$

$$x \in \mathbb{R}^D, \phi(x) \in \mathbb{R}^C, C > D!$$

Feature Transformation

Key Idea: Transform inputs x using a feature map $\phi(x)$

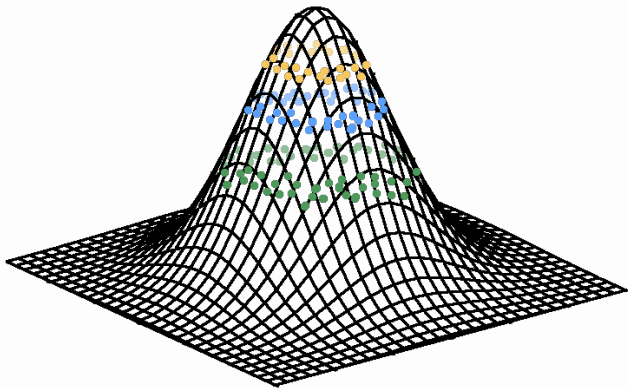
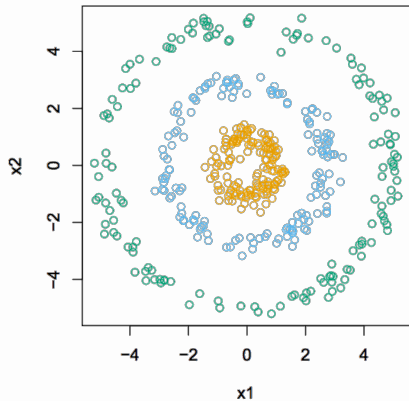
$$x \in \mathbb{R}^D, \phi(x) \in \mathbb{R}^C, C > D!$$



Feature Transformation

Key Idea: Transform inputs x using a feature map $\phi(x)$

$$x \in \mathbb{R}^D, \phi(x) \in \mathbb{R}^C, C > D!$$



Kernel PCA

$$X = [\mathbf{x}_1; \dots; \mathbf{x}_N],$$

$$\kappa(\mathbf{x}_i, \mathbf{x}_j) = \phi(\mathbf{x}_i)^\top \phi(\mathbf{x}_j)$$

(kernel function)

$$\Phi(X) = [\phi(\mathbf{x}_1); \dots; \phi(\mathbf{x}_N)]$$

$$K = \Phi(X)^\top \Phi(X)$$

(kernel matrix)

Kernel PCA

$$X = [\mathbf{x}_1; \dots; \mathbf{x}_N],$$

$$\kappa(\mathbf{x}_i, \mathbf{x}_j) = \phi(\mathbf{x}_i)^\top \phi(\mathbf{x}_j)$$

(kernel function)

$$\Phi(X) = [\phi(\mathbf{x}_1); \dots; \phi(\mathbf{x}_N)]$$

$$K = \Phi(X)^\top \Phi(X)$$

(kernel matrix)

$$S = \frac{1}{N} X X^\top$$

$$S\mathbf{v} = \lambda\mathbf{v}$$

Kernel PCA

$$X = [\mathbf{x}_1; \dots; \mathbf{x}_N],$$

$$\kappa(\mathbf{x}_i, \mathbf{x}_j) = \phi(\mathbf{x}_i)^\top \phi(\mathbf{x}_j)$$

(kernel function)

$$\Phi(X) = [\phi(\mathbf{x}_1); \dots; \phi(\mathbf{x}_N)]$$

$$K = \Phi(X)^\top \Phi(X)$$

(kernel matrix)

$$S = \frac{1}{N} X X^\top$$

$$S\mathbf{v} = \lambda\mathbf{v}$$

$$S = \frac{1}{N} \Phi(X) \Phi(X)^\top$$

$$\mathbf{v} = \Phi(X) \mathbf{a} = \sum_{i=1}^N a_i \phi(\mathbf{x}_i)$$

(...some algebra)

$$K\mathbf{a} = N\lambda\mathbf{a}$$

Kernel PCA

$$X = [\mathbf{x}_1; \dots; \mathbf{x}_N],$$

$$\kappa(\mathbf{x}_i, \mathbf{x}_j) = \phi(\mathbf{x}_i)^\top \phi(\mathbf{x}_j)$$

(kernel function)

$$\Phi(X) = [\phi(\mathbf{x}_1); \dots; \phi(\mathbf{x}_N)]$$

$$K = \Phi(X)^\top \Phi(X)$$

(kernel matrix)

$$S = \frac{1}{N} X X^\top$$

$$S \mathbf{v} = \lambda \mathbf{v}$$

$$S = \frac{1}{N} \Phi(X) \Phi(X)^\top$$

$$\mathbf{v} = \Phi(X) \mathbf{a} = \sum_{i=1}^N a_i \phi(\mathbf{x}_i)$$

(...some algebra)

$$K \mathbf{a} = N \lambda \mathbf{a}$$

Kernel Trick

No need to compute $\phi(\mathbf{x})$ —only $\kappa(\mathbf{x}_i, \mathbf{x}_j)$!

Kernel PCA: Example

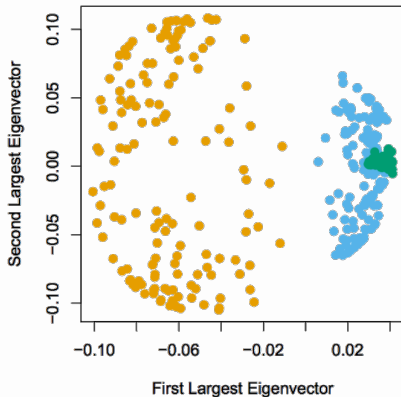
$$\kappa(\mathbf{x}_i, \mathbf{x}_j) = \exp\left(\frac{-\|\mathbf{x}_i - \mathbf{x}_j\|^2}{2\sigma^2}\right)$$

(RBF Kernel)

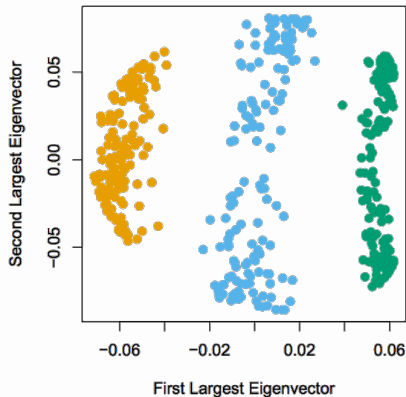
Kernel PCA: Example

$$\kappa(\mathbf{x}_i, \mathbf{x}_j) = \exp\left(\frac{-\|\mathbf{x}_i - \mathbf{x}_j\|^2}{2\sigma^2}\right)$$

(RBF Kernel)



RBF ($\sigma^2 = 0.5$)



RBF ($\sigma^2 = 0.01$)

Kernel PCA

Choosing the right kernel

- Not an easy choice

Kernel PCA

Choosing the right kernel

- Not an easy choice
- Construct by hand/eye where feasible

Kernel PCA

Choosing the right kernel

- Not an easy choice
- Construct by hand/eye where feasible
- Possible to learn kernel matrices K directly from data!

Kernel PCA

Choosing the right kernel

- Not an easy choice
- Construct by hand/eye where feasible
- Possible to learn kernel matrices K directly from data!

Kernel PCA

Choosing the right kernel

- Not an easy choice
- Construct by hand/eye where feasible
- Possible to learn kernel matrices K directly from data!

Several non-linear dimensionality reduction methods can be viewed as kernel PCA, with kernels learned from data [1]

1. J. Ham et al, A Kernel View of the Dimensionality Reduction of Manifolds, 2004

Visualisation

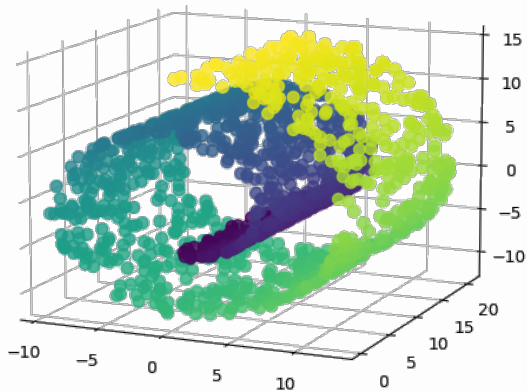
Manifold Hypothesis

High-dimensional data in the real world really lies on low-dimensional manifolds within that high-dimensional space.



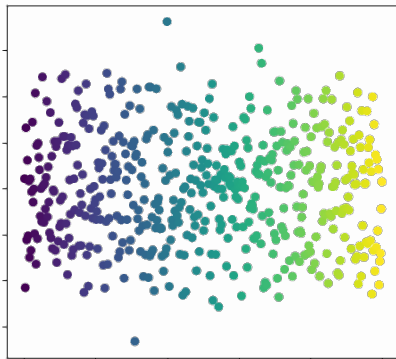
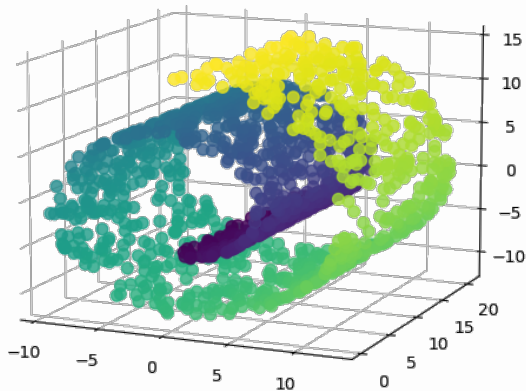
Manifold Hypothesis

High-dimensional data in the real world really lies on low-dimensional manifolds within that high-dimensional space.



Manifold Hypothesis

High-dimensional data in the real world really lies on low-dimensional manifolds within that high-dimensional space.



Overview

Key Ideas

- Difficult to construct a single global transformation of data
- Focus instead on some *local* measure of closeness
- Project data x onto lower-dimensional manifold as e
- **Question:** can we *preserve* local measure of closeness?

Overview

Key Ideas

- Difficult to construct a single global transformation of data
- Focus instead on some *local* measure of closeness
- Project data x onto lower-dimensional manifold as e
- **Question:** can we *preserve* local measure of closeness?

Let $\mathcal{X}$ denote the high-dimensional data space, and $\mathcal{E}$ denote the low-dimensional manifold space. We can define the following distance measures on the two spaces respectively

$$\mathcal{D}_{\mathcal{X}}(x_i, x_j)$$

$$\mathcal{D}_{\mathcal{E}}(e_i, e_j)$$

Multi-Dimensional Scaling (MDS)

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between *every pair of samples in original data X*

Multi-Dimensional Scaling (MDS)

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between *every pair of samples in original data X*

$$\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) = \|\mathbf{x}_i - \mathbf{x}_j\| \qquad \mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_j) = \|\mathbf{e}_i - \mathbf{e}_j\| \qquad \text{(example)}$$

Multi-Dimensional Scaling (MDS)

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between *every pair of samples in original data X*

$$\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) = \|\mathbf{x}_i - \mathbf{x}_j\| \quad \mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_j) = \|\mathbf{e}_i - \mathbf{e}_j\| \quad (\text{example})$$

$$\text{Objective: } \min \sum_{i,j} (\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) - \mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_j))^2$$

Multi-Dimensional Scaling (MDS)

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between *every pair of samples in original data X*

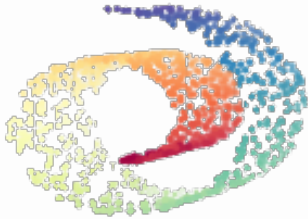
$$\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) = \|\mathbf{x}_i - \mathbf{x}_j\| \quad \mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_j) = \|\mathbf{e}_i - \mathbf{e}_j\| \quad (\text{example})$$

$$\text{Objective: } \min \sum_{i,j} (\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) - \mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_j))^2$$

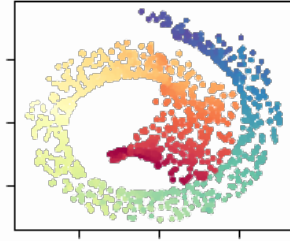
- distance metrics $\mathcal{D}_{\mathcal{X}}, \mathcal{D}_{\mathcal{E}}$ could really be anything
- choosing L_2 helps make optimisation simpler

MDS: Example

Swiss Roll



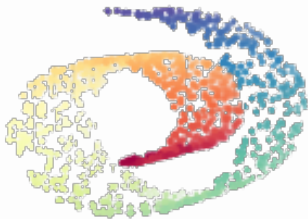
Data



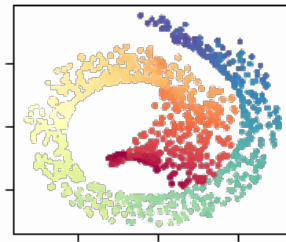
MDS (Euclidean)

MDS: Example

Swiss Roll



Data



MDS (Euclidean)

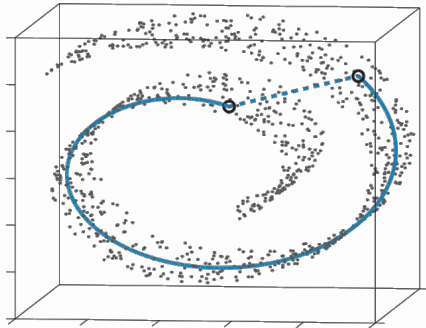
- Projects down to 2D while preserving distances
- Preserving distances *outside* the manifold!

Isomap

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between points on the *embedded manifold*, not arbitrary distance!

Isomap

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between points on the *embedded manifold*, not arbitrary distance!

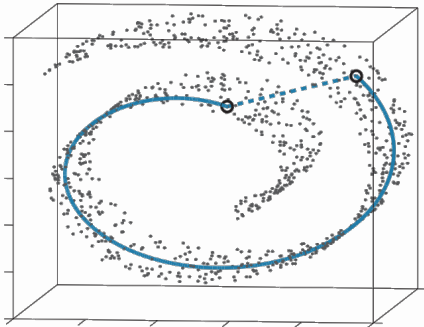


Isomap

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between points on the *embedded manifold*, not arbitrary distance!

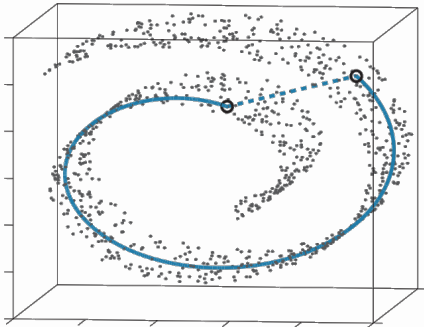
Find the Geodesic

- Generate nearest-neighbour graph $\mathcal{G}$ on data



Isomap

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between points on the *embedded manifold*, not arbitrary distance!

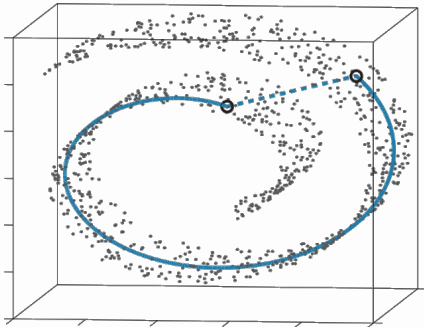


Find the Geodesic

- Generate nearest-neighbour graph $\mathcal{G}$ on data
- Shortest distance between points in this graph
 - Floyd-Warshall algorithm (*all pairs* shortest path)

Isomap

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between points on the *embedded manifold*, not arbitrary distance!

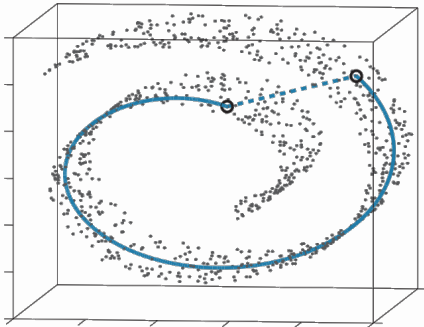


Find the Geodesic

- Generate nearest-neighbour graph $\mathcal{G}$ on data
- Shortest distance between points in this graph
 - Floyd-Warshall algorithm (*all pairs* shortest path)
- Perform MDS with $\mathcal{D}_{\mathcal{X}}(x_i, x_j) = \mathcal{G}_{FW}(x_i, x_j)$

Isomap

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between points on the *embedded manifold*, not arbitrary distance!

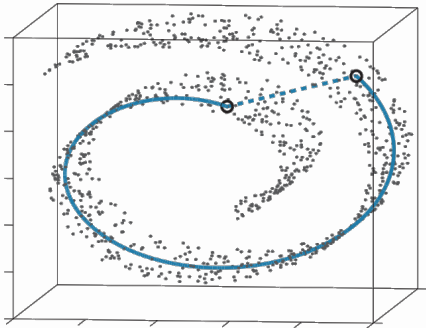


Find the Geodesic

- Generate nearest-neighbour graph $\mathcal{G}$ on data
- Shortest distance between points in this graph
 - Floyd-Warshall algorithm (*all pairs* shortest path)
- Perform MDS with $\mathcal{D}_{\mathcal{X}}(x_i, x_j) = \mathcal{G}_{\text{FW}}(x_i, x_j)$

Isomap

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the distance between points on the *embedded manifold*, not arbitrary distance!

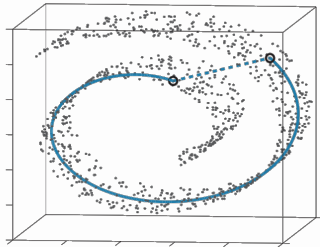


Find the Geodesic

- Generate nearest-neighbour graph $\mathcal{G}$ on data
- Shortest distance between points in this graph
 - Floyd-Warshall algorithm (*all pairs* shortest path)
- Perform MDS with $\mathcal{D}_{\mathcal{X}}(x_i, x_j) = \mathcal{G}_{\text{FW}}(x_i, x_j)$

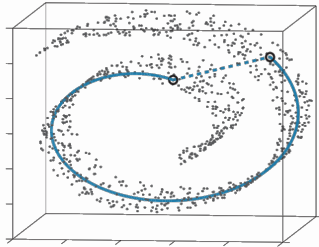
$$\text{Objective: } \min \sum_{i,j} (\mathcal{G}_{\text{FW}}(x_i, x_j) - \|e_i - e_j\|)^2$$

Isomap: Manifold

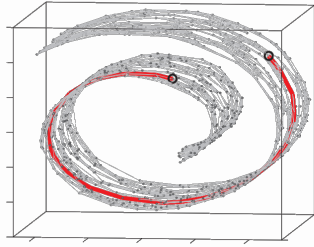


Data

Isomap: Manifold



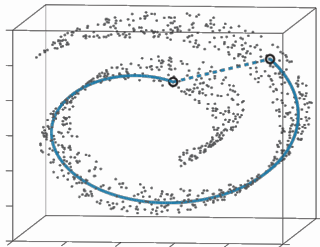
Data



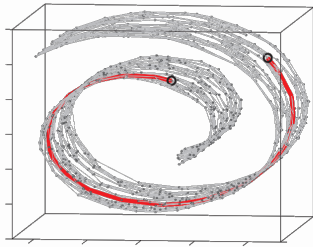
Graph

+ Floyd-Warshall

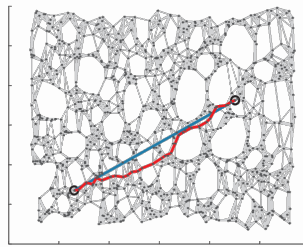
Isomap: Manifold



Data

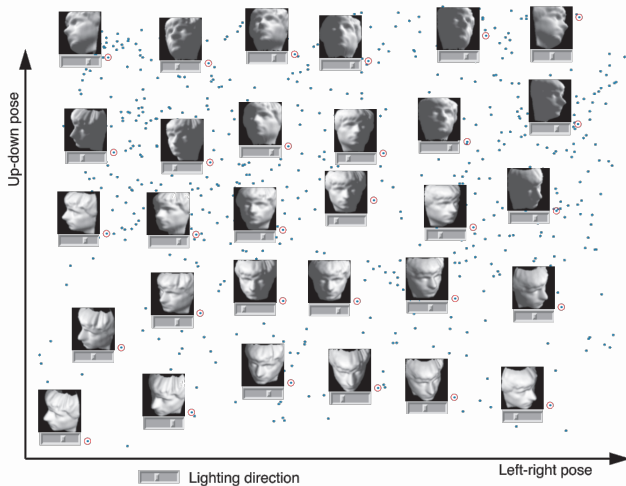


Graph
+ Floyd-Warshall

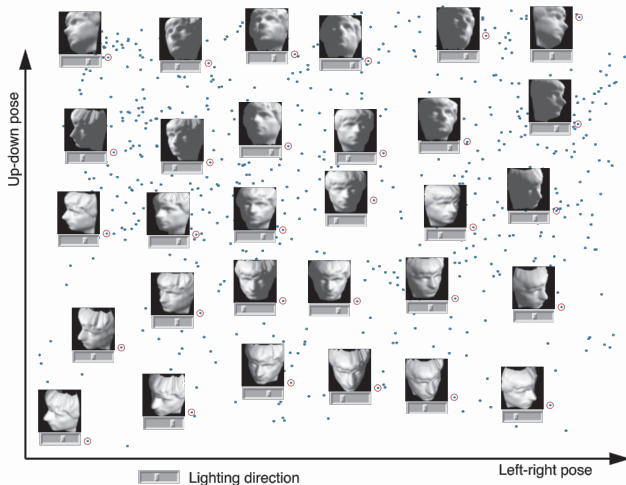


MDS

Isomap: Example and Issues



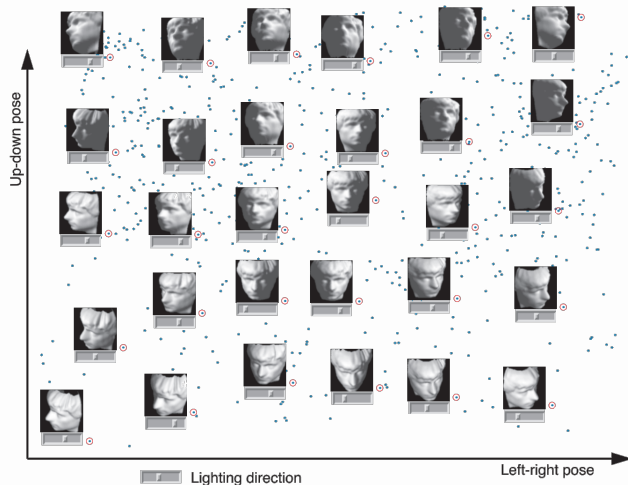
Isomap: Example and Issues



Issues

- requires *uniform* and *dense* sampling on manifold

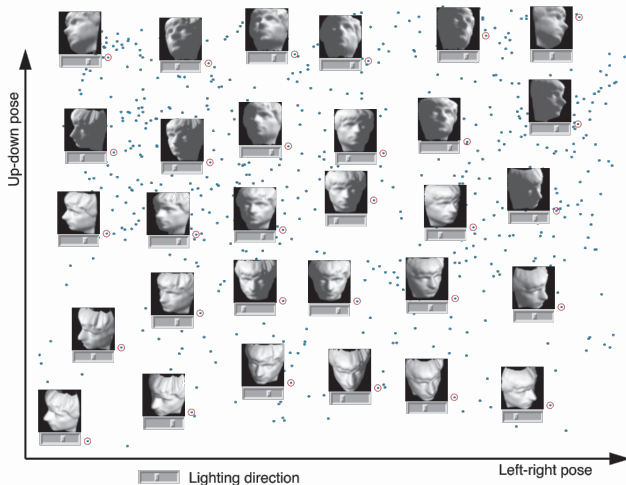
Isomap: Example and Issues



Issues

- requires *uniform* and *dense* sampling on manifold
- prone to topological instabilities (effect of noise, non-convexity)

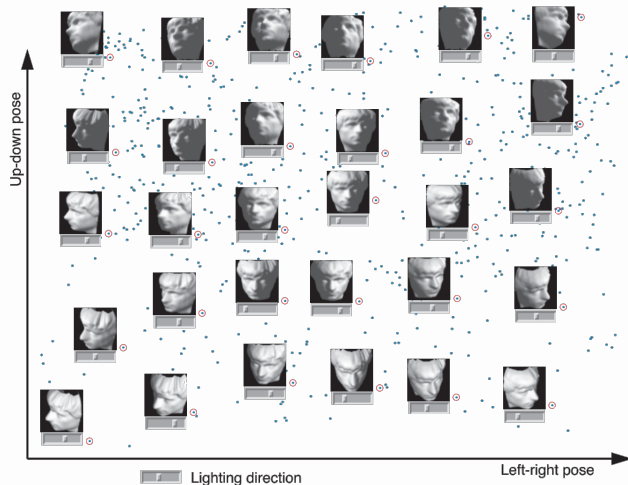
Isomap: Example and Issues



Issues

- requires *uniform* and *dense* sampling on manifold
- prone to topological instabilities (effect of noise, non-convexity)
- can get disconnected graphs!

Isomap: Example and Issues



Issues

- requires *uniform* and *dense* sampling on manifold
- prone to topological instabilities (effect of noise, non-convexity)
- can get disconnected graphs!
- slow with size of data — Floyd-Warshall is $O(N^3)$!

Locally Linear Embeddings (LLE)

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the *linear transform that reconstructs points* from the K nearest neighbours

Locally Linear Embeddings (LLE)

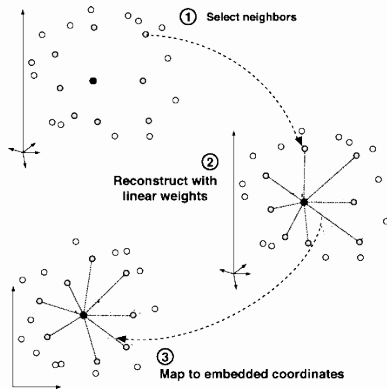
Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the *linear transform that reconstructs points from the K nearest neighbours*

$$\arg \min_{w_1, \dots, w_K} \sum_{i=1}^N (\mathbf{x}_i - \sum_{\mathbf{x}_k \in n(\mathbf{x}_i)} w_k \mathbf{x}_k)^2$$

$$\arg \min_{e_1, \dots, e_N} \sum_{i=1}^N (e_i - \sum_{e_k \in n(e_i)} w_k e_k)^2$$

Algorithm

1. Find the weights w_k
2. Fix weights w_k and find optimal embeddings e_i
solved using a (sparse) eigen-decomposition



Locally Linear Embeddings (LLE)

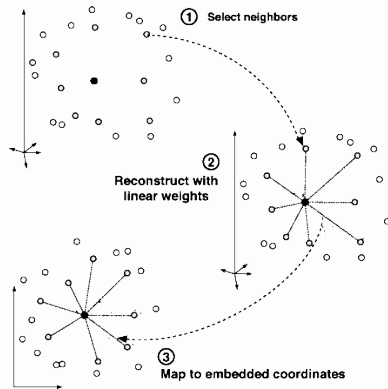
Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the *linear transform that reconstructs points from the K nearest neighbours*

$$\arg \min_{w_1, \dots, w_K} \sum_{i=1}^N (\mathbf{x}_i - \sum_{\mathbf{x}_k \in n(\mathbf{x}_i)} w_k \mathbf{x}_k)^2$$

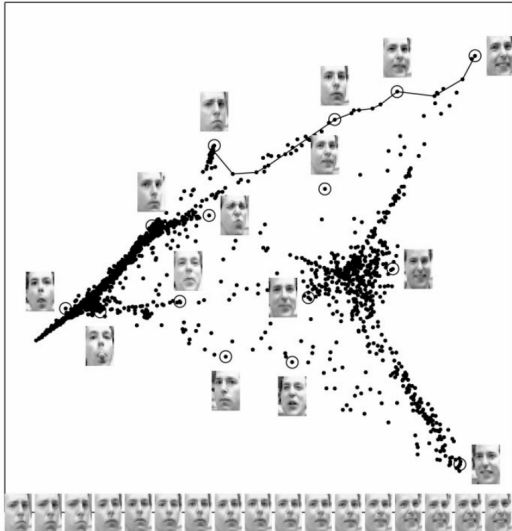
$$\arg \min_{e_1, \dots, e_N} \sum_{i=1}^N (e_i - \sum_{e_k \in n(e_i)} w_k e_k)^2$$

Algorithm

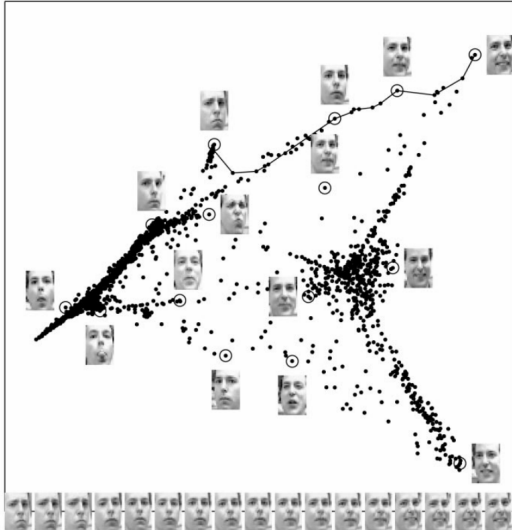
1. Find the weights w_k
2. Fix weights w_k and find optimal embeddings e_i
solved using a (sparse) eigen-decomposition



LLE: Example and Issues



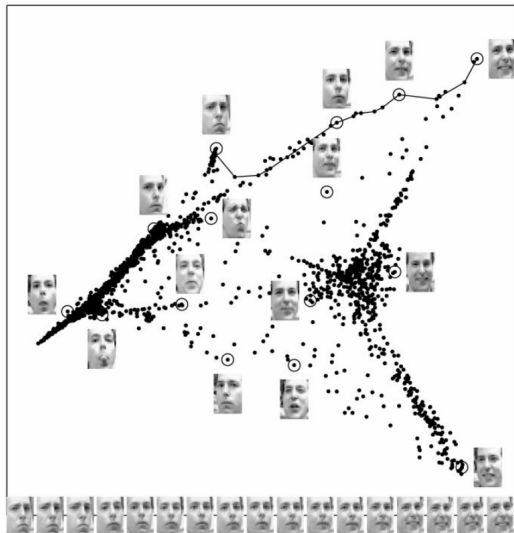
LLE: Example and Issues



Advantages

- globally optimal

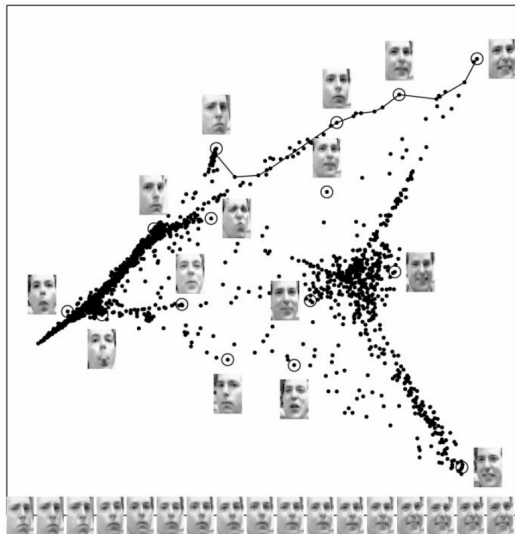
LLE: Example and Issues



Advantages

- globally optimal
- only bottleneck—finding e_i

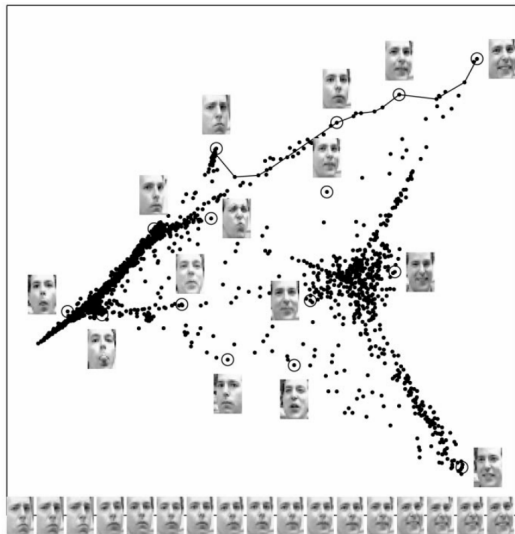
LLE: Example and Issues



Advantages

- globally optimal
- only bottleneck—finding e_i
- *not* preserving specific distance

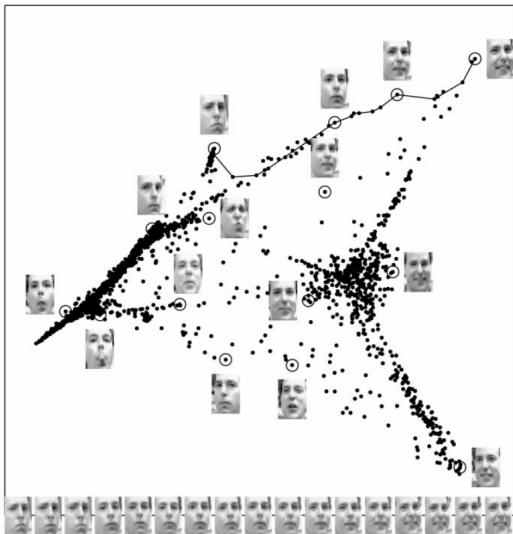
LLE: Example and Issues



Advantages

- globally optimal
- only bottleneck—finding e_i
- *not* preserving specific distance

LLE: Example and Issues



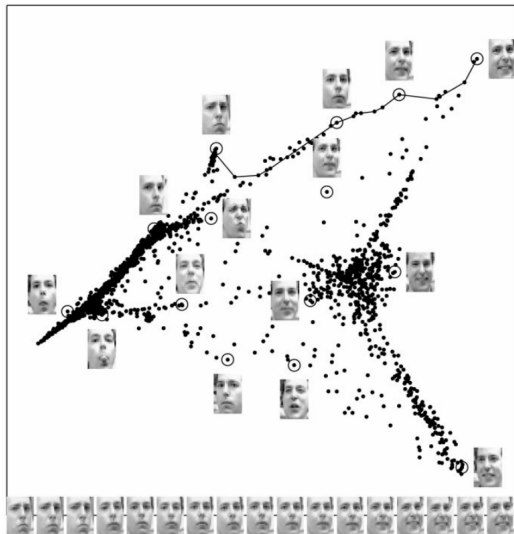
Advantages

- globally optimal
- only bottleneck—finding e_i
- *not* preserving specific distance

Issues

- reliance on local smoothness

LLE: Example and Issues



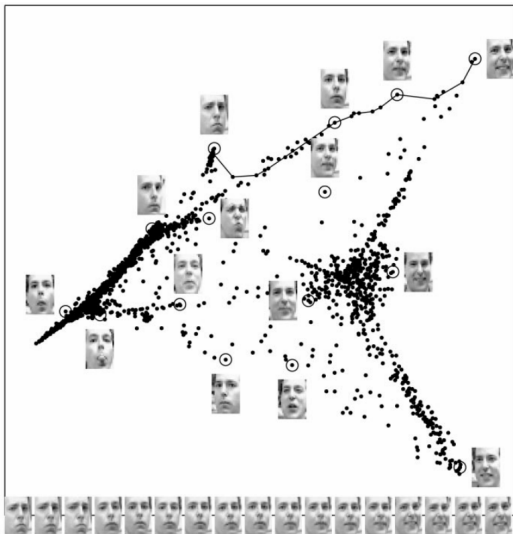
Advantages

- globally optimal
- only bottleneck—finding e_i
- *not* preserving specific distance

Issues

- reliance on local smoothness
- sensitive to noise

LLE: Example and Issues



Advantages

- globally optimal
- only bottleneck—finding e_i
- *not* preserving specific distance

Issues

- reliance on local smoothness
- sensitive to noise
- no theoretical guarantees about manifold

t -distributed Stochastic Neighbour Embedding (t -SNE)

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the *probability distribution* over *pairwise similarities* between points in the space.

t -distributed Stochastic Neighbour Embedding (t -SNE)

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the *probability distribution over pairwise similarities* between points in the space.

$$p(\mathbf{x}_j|\mathbf{x}_i) = \frac{\exp(-\|\mathbf{x}_i - \mathbf{x}_j\|^2/2\sigma_i^2)}{\sum_{k \neq i} \exp(-\|\mathbf{x}_i - \mathbf{x}_k\|^2/2\sigma_i^2)} \quad (i \neq j)$$

$$p(\mathbf{x}_i|\mathbf{x}_i) = 0 \quad \sum_i p(\mathbf{x}_j|\mathbf{x}_i) = 1$$

$$\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) = \frac{1}{2N} (p(\mathbf{x}_i|\mathbf{x}_j) + p(\mathbf{x}_j|\mathbf{x}_i))$$

$$\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_i) = 0 \quad \sum_{ij} \mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) = 1$$

t -distributed Stochastic Neighbour Embedding (t -SNE)

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the *probability distribution over pairwise similarities* between points in the space.

$$p(\mathbf{x}_j|\mathbf{x}_i) = \frac{\exp(-\|\mathbf{x}_i - \mathbf{x}_j\|^2/2\sigma_i^2)}{\sum_{k \neq i} \exp(-\|\mathbf{x}_i - \mathbf{x}_k\|^2/2\sigma_i^2)} \quad (i \neq j)$$

$$p(\mathbf{x}_i|\mathbf{x}_i) = 0 \quad \sum_i p(\mathbf{x}_j|\mathbf{x}_i) = 1$$

$$\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) = \frac{1}{2N} (p(\mathbf{x}_i|\mathbf{x}_j) + p(\mathbf{x}_j|\mathbf{x}_i))$$

$$\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_i) = 0 \quad \sum_{ij} \mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) = 1$$

$$\mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_j) := \text{Student-}t(\|\mathbf{e}_i - \mathbf{e}_j\|^2, \nu = 1)$$

$$\mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_j) = \frac{(1 + \|\mathbf{e}_i - \mathbf{e}_j\|^2)^{-1}}{\sum_k \sum_{l \neq k} (1 + \|\mathbf{e}_k - \mathbf{e}_l\|^2)^{-1}}$$

$$\mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_i) = 0$$

Student- t allows dissimilar points in $\mathcal{X}$ to be modelled far away in $\mathcal{E}$!

t -distributed Stochastic Neighbour Embedding (t -SNE)

Project data from $\mathcal{X}$ to $\mathcal{E}$ while preserving the *probability distribution over pairwise similarities* between points in the space.

$$p(\mathbf{x}_j|\mathbf{x}_i) = \frac{\exp(-\|\mathbf{x}_i - \mathbf{x}_j\|^2/2\sigma_i^2)}{\sum_{k \neq i} \exp(-\|\mathbf{x}_i - \mathbf{x}_k\|^2/2\sigma_i^2)} \quad (i \neq j)$$

$$p(\mathbf{x}_i|\mathbf{x}_i) = 0 \quad \sum_i p(\mathbf{x}_j|\mathbf{x}_i) = 1$$

$$\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) = \frac{1}{2N} (p(\mathbf{x}_i|\mathbf{x}_j) + p(\mathbf{x}_j|\mathbf{x}_i))$$

$$\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_i) = 0 \quad \sum_{ij} \mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) = 1$$

$$\mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_j) := \text{Student-}t(\|\mathbf{e}_i - \mathbf{e}_j\|^2, \nu = 1)$$

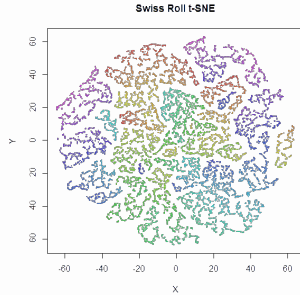
$$\mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_j) = \frac{(1 + \|\mathbf{e}_i - \mathbf{e}_j\|^2)^{-1}}{\sum_k \sum_{l \neq k} (1 + \|\mathbf{e}_k - \mathbf{e}_l\|^2)^{-1}}$$

$$\mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_i) = 0$$

Student- t allows dissimilar points in $\mathcal{X}$ to be modelled far away in $\mathcal{E}$!

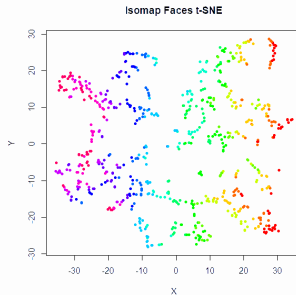
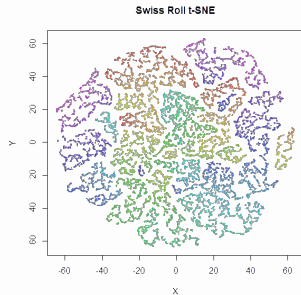
$$\text{Objective: } \min \text{KL}(\mathcal{D}_{\mathcal{X}}(\mathbf{x}_i, \mathbf{x}_j) \| \mathcal{D}_{\mathcal{E}}(\mathbf{e}_i, \mathbf{e}_j))$$

t -SNE: Examples



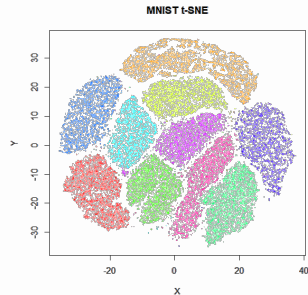
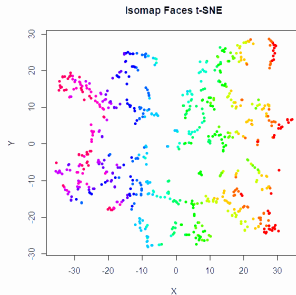
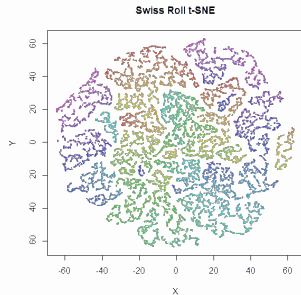
Figures: <https://jlmelville.github.io/uwot/umap-examples.html>

t -SNE: Examples



Figures: <https://jlmelville.github.io/uwot/umap-examples.html>

t -SNE: Examples



Figures: <https://jlmelville.github.io/uwot/umap-examples.html>

So Far

- Different (relatively simple) aspects may be preserved in mapping $\mathcal{X}$ to $\mathcal{E}$

So Far

- Different (relatively simple) aspects may be preserved in mapping $\mathcal{X}$ to $\mathcal{E}$
- Trade-off in terms of simplicity of assumption and ease of optimisation

So Far

- Different (relatively simple) aspects may be preserved in mapping $\mathcal{X}$ to $\mathcal{E}$
- Trade-off in terms of simplicity of assumption and ease of optimisation
- t -SNE quite popular
 - can incorporate both global and local structure (distributional)
 - works on large scale *and* high-dimensional data

So Far

- Different (relatively simple) aspects may be preserved in mapping $\mathcal{X}$ to $\mathcal{E}$
- Trade-off in terms of simplicity of assumption and ease of optimisation
- t -SNE quite popular
 - can incorporate both global and local structure (distributional)
 - works on large scale *and* high-dimensional data

So Far

- Different (relatively simple) aspects may be preserved in mapping $\mathcal{X}$ to $\mathcal{E}$
- Trade-off in terms of simplicity of assumption and ease of optimisation
- t -SNE quite popular
 - can incorporate both global and local structure (distributional)
 - works on large scale *and* high-dimensional data

Issue

None of these methods learn an *explicit* metric

- $\{x_1, \dots, x_N\} \xrightarrow{f} \{e_1, \dots, e_N\}$ but $x_t \not\xrightarrow{\quad} e_t$ for unseen x_t
- cannot project an unseen point without redoing the optimisation!

Uniform Manifold Approximation & Projection (UMAP)

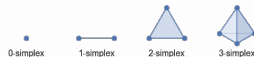
Leverage Riemannian geometry and topology to construct a general framework for manifold learning and dimensionality reduction.

Uniform Manifold Approximation & Projection (UMAP)

Leverage Riemannian geometry and topology to construct a general framework for manifold learning and dimensionality reduction.

Overview

- Use simplices as basic building block

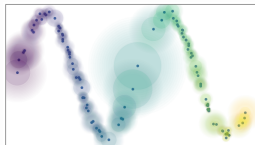
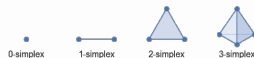


Uniform Manifold Approximation & Projection (UMAP)

Leverage Riemannian geometry and topology to construct a general framework for manifold learning and dimensionality reduction.

Overview

- Use simplices as basic building block
- Estimate connectivity and distances between points

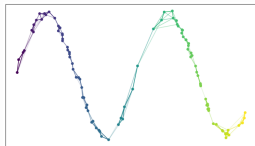
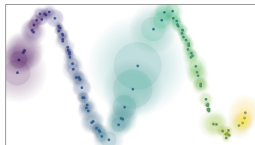
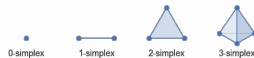


Uniform Manifold Approximation & Projection (UMAP)

Leverage Riemannian geometry and topology to construct a general framework for manifold learning and dimensionality reduction.

Overview

- Use simplices as basic building block
- Estimate connectivity and distances between points
- Construct graph that captures topology of manifold!



L. McInnes et al, UMAP: Uniform Manifold Approximation and Projection for Dimension Reduction, 2018

UMAP: Dimensionality Reduction

- Construct faithful topological representation of data

UMAP: Dimensionality Reduction

- Construct faithful topological representation of data
- Compute *cross-entropy* between topological structures of $\mathcal{X}$ and $\mathcal{E}$ in terms of the simplices (building blocks)

UMAP: Dimensionality Reduction

- Construct faithful topological representation of data
- Compute *cross-entropy* between topological structures of $\mathcal{X}$ and $\mathcal{E}$ in terms of the simplices (building blocks)
- Optimise low-dimensional representation to have minimum cross-entropy

UMAP: Dimensionality Reduction

- Construct faithful topological representation of data
- Compute *cross-entropy* between topological structures of $\mathcal{X}$ and $\mathcal{E}$ in terms of the simplices (building blocks)
- Optimise low-dimensional representation to have minimum cross-entropy

UMAP: Dimensionality Reduction

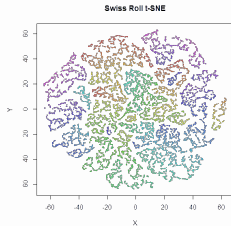
- Construct faithful topological representation of data
- Compute *cross-entropy* between topological structures of $\mathcal{X}$ and $\mathcal{E}$ in terms of the simplices (building blocks)
- Optimise low-dimensional representation to have minimum cross-entropy

Advantages

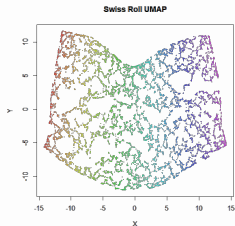
- Can learn a metric, so computing $x_t \xrightarrow{f} e_t$ for unseen x_t is feasible (when labels available)!
- Is fast and scalable
- Can be run unsupervised, supervised, or even weakly supervised!

UMAP: Examples and Comparisons

t -SNE



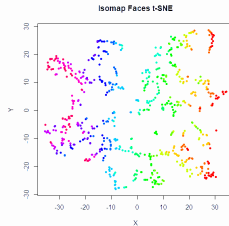
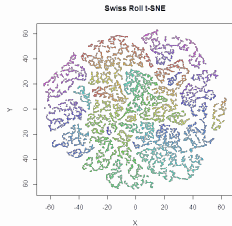
UMAP



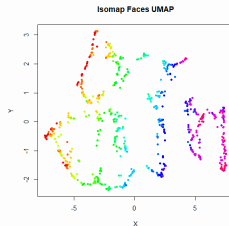
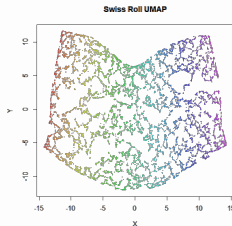
Figures: <https://jlmelville.github.io/uwot/umap-examples.html>

UMAP: Examples and Comparisons

t -SNE



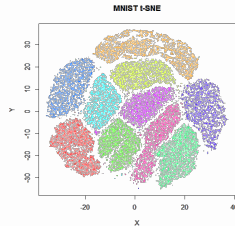
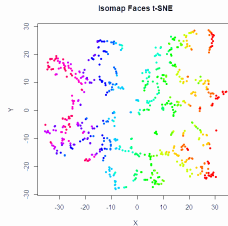
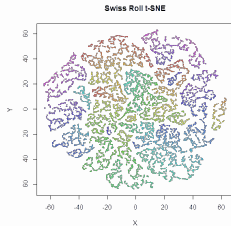
UMAP



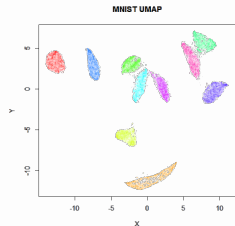
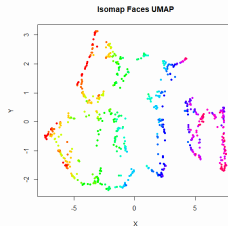
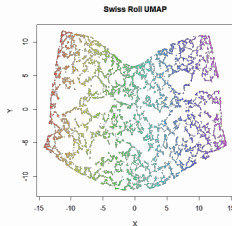
Figures: <https://jlmelville.github.io/uwot/umap-examples.html>

UMAP: Examples and Comparisons

t-SNE



UMAP



Figures: <https://jlmelville.github.io/uwot/umap-examples.html>

Summary

- Non-linear dimensionality reduction helps visualise complex data in low dimensions—relies on the manifold hypothesis

Summary

- Non-linear dimensionality reduction helps visualise complex data in low dimensions—relies on the manifold hypothesis
- Can explore ways to transform data to run linear versions of algorithms (e.g. kernel PCA)

Summary

- Non-linear dimensionality reduction helps visualise complex data in low dimensions—relies on the manifold hypothesis
- Can explore ways to transform data to run linear versions of algorithms (e.g. kernel PCA)
- Most methods exploit the nearest neighbour graph in some form (e.g. MDS, Isomap, LLE, etc.)

Summary

- Non-linear dimensionality reduction helps visualise complex data in low dimensions—relies on the manifold hypothesis
- Can explore ways to transform data to run linear versions of algorithms (e.g. kernel PCA)
- Most methods exploit the nearest neighbour graph in some form (e.g. MDS, Isomap, LLE, etc.)
- Data is typically required to be clean (not much noise) and dense
...not too strong a requirement, especially with vision or language

Summary

- Non-linear dimensionality reduction helps visualise complex data in low dimensions—relies on the manifold hypothesis
- Can explore ways to transform data to run linear versions of algorithms (e.g. kernel PCA)
- Most methods exploit the nearest neighbour graph in some form (e.g. MDS, Isomap, LLE, etc.)
- Data is typically required to be clean (not much noise) and dense ...not too strong a requirement, especially with vision or language
- Many approaches to choose from— t -SNE and UMAP most popular